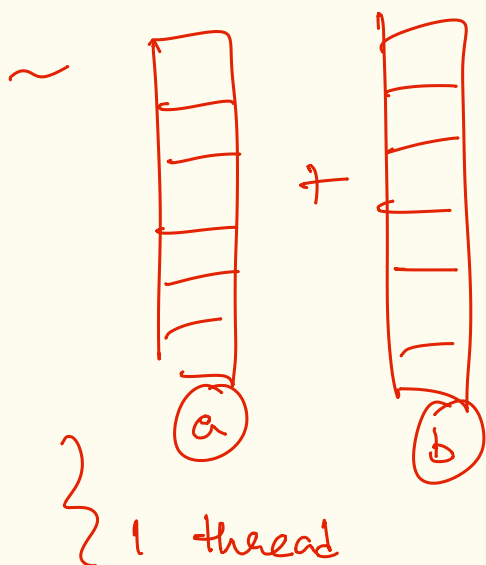


CS695

introduction to GPUs
& GPU virtualization

~ GPUs

(i) SIMD
Single instruction multiple tasks data.
logic



loop {
 $a[i], b[i]$
}

Quiz 2

Fri. 11/4

8:30 am

Faal platform
hands-on

A4

~ assign a paper
reading
material

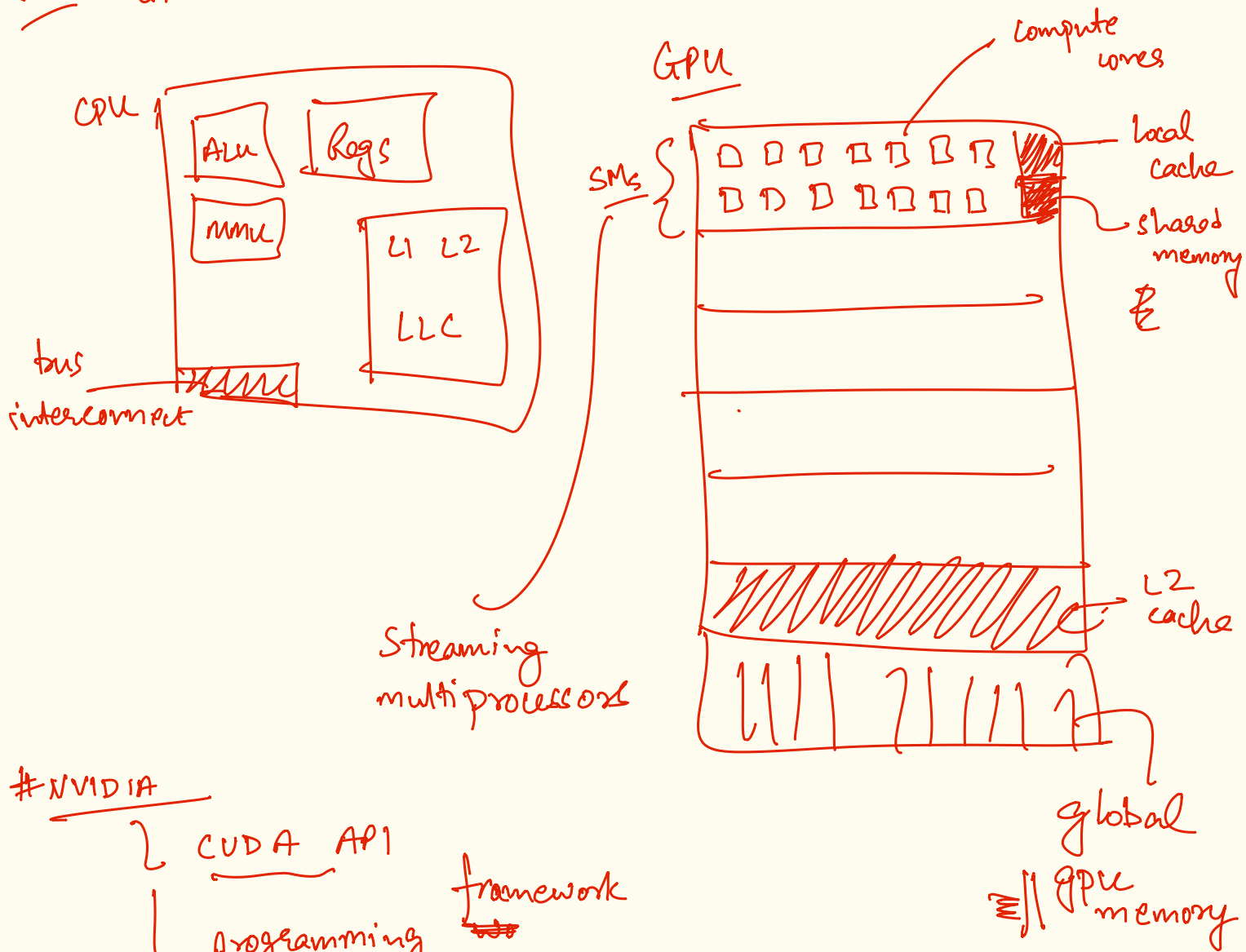
~ submit 5-6

slides of content
on the paper.

} } } }

$a[i] + b[i]$ ~ same operation
on all threads
thread specific of execution.

(ii) GPU architecture



NVIDIA

CUDA API
programming framework

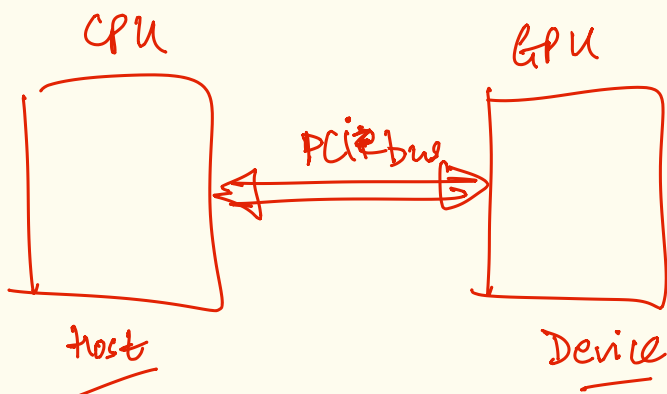
threads — instance of parallel execution
blocks — sets of threads — scheduled on a core
grids — schedulable granularity on an SM.
blocks of your program

CUDA programs

tl/w. discrete GPAs (i)

integrated / HSA GPUs. (ii)

(i)



(ii)

GPU & CPU share a memory bus

⇒ both can same access phy. addresses

copy mem from H to D

do GPU work ——— NVIDIA cuda kernel

copy mem from D to H

anatomy cuda program

-- global -- add, (int* in, int* out)

{

tid = (blockid + threadid) / blocksizes)

a[tid] = a[tid] + b[tid];

}

main () {

int a[], b[];

ga = cudaMalloc ();

gb = cudaMalloc ();

copy mem H to D (a, ga)

→ << #blocks, #threads per block >>> add (ga, gb); (b, gb);

}

<< 2, 2 >>>
2 2

block 0 tid 0
1

block 1 tid 0
1

✓ << 2, 2 >>>

⊗ << 1, 4 >>>

* applⁿs.

vector-based arithmetic

- + AI/ML
- + image processing
- + audio/video analysis
- + cryptography

virtualization

- + sharing
- + multiplexing
- + utilization
- + performance

GPU virtualization

- + MIG } Multi-Instance GPU
- + vGPU } vGPU
- + MPS ~ ^{gpu}slur ~ driver provided per process GPU resource
- * + API remoting / virtualization
- * + Kernel slicing
- + FaaS

are at a higher imp. level for GPUs.

hw assisted & GPU driver
breakup 1 GPU into 2 GPUs
(device interface)

MIG - spatial multiplexing
vGPU - temporal multiplexing

$$\frac{4}{7} \quad \& \quad \frac{3}{7}$$

all GPU resources
divided spatially
in this ratio

- memory is divided spatially

- cores divided / multiplexed temporally

+ best-effort ~ if vGPU has no work schedule from other vGPUs.

